

**MIT 9.520/6.7910,
Statistical Learning Theory and Applications**

Class 03: Least squares and overpamaterization

Lorenzo Rosasco

Learning problem and algorithms

Solve

$$\min_{f \in \mathcal{F}} L(f) = L(f_P), \quad L(f) = \mathbb{E}_{(x,y) \sim P}[\ell(y, f(x))],$$

given only

$$S_n = (x_1, y_1), \dots, (x_n, y_n) \sim P^n.$$

Learning problem and algorithms

Solve

$$\min_{f \in \mathcal{F}} L(f) = L(f_P), \quad L(f) = \mathbb{E}_{(x,y) \sim P}[\ell(y, f(x))],$$

given only

$$S_n = (x_1, y_1), \dots, (x_n, y_n) \sim P^n.$$

Learning algorithm

$$S_n \rightarrow \hat{f} = \hat{f}_{S_n},$$

$\hat{f}$ estimates f_P given the observed examples S_n .

Learning problem and algorithms

Solve

$$\min_{f \in \mathcal{F}} L(f) = L(f_P), \quad L(f) = \mathbb{E}_{(x,y) \sim P}[\ell(y, f(x))],$$

given only

$$S_n = (x_1, y_1), \dots, (x_n, y_n) \sim P^n.$$

Learning algorithm

$$S_n \rightarrow \hat{f} = \hat{f}_{S_n},$$

$\hat{f}$ estimates f_P given the observed examples S_n .

How can we design a learning algorithm?

Classic algorithm design: representation

Pick a model, e.g. parameterized

Linearly

$$w^T x \mapsto w^T \Phi(x)$$

Classic algorithm design: representation

Pick a model, e.g. parameterized

Linearly

$$w^T x \mapsto w^T \Phi(x)$$

or

Non linearly

$$w^T \sigma(Wx) \mapsto w^T \sigma(W_L \dots \sigma(W_2 \sigma(W_1 x))$$

Classic algorithm design: representation

Pick a model, e.g. parameterized

Linearly

$$w^T x \mapsto w^T \Phi(x)$$

or

Non linearly

$$w^T \sigma(Wx) \mapsto w^T \sigma(W_L \dots \sigma(W_2 \sigma(W_1 x))$$

This includes various choices like features/kernels, architecture and possible constraints e.g.

$$\|w\| \leq R, \quad \|W_j\| \leq R_j.$$

Classic algorithm design: optimization

Given some model, e.g. $f_w(x)$, $\|w\| \leq R$, consider the empirical risk minimization (ERM)

$$\min_{\|w\| \leq R} \hat{L}(f_w) \quad \hat{L}(f_w) = \frac{1}{n} \sum_{i=1}^n \ell(y_i, f_w(x_i)),$$

or

$$\min_w \hat{L}(f_w) + \lambda \|w\|.$$

Outline

Overparameterization and bias

Regularization and stability

Representer theorem and numerical complexity

Beyond ℓ^2 regularization

Linear functions and linear spaces

Let $\mathcal{H}$ be the space of linear functions

$$f(x) = w^\top x.$$

Then,

- ▶ $f \leftrightarrow w$ is one to one,
- ▶ inner product $\langle f, \bar{f} \rangle_{\mathcal{H}} := w^\top \bar{w}$,
- ▶ norm/metric $\|f - \bar{f}\|_{\mathcal{H}} := \|w - \bar{w}\|$.

Linear functions are the conceptual building block of most functions class.

Linear least squares

ERM with least squares also called ordinary least squares (OLS)

$$\min_{w \in \mathbb{R}^d} \underbrace{\frac{1}{n} \sum_{i=1}^n (y_i - w^\top x_i)^2}_{\hat{L}(w)}.$$

Discussing computations provides many insights.

Matrices and linear systems

Let $\hat{X} \in \mathbb{R}^{nd}$ and $\hat{Y} \in \mathbb{R}^n$. Then

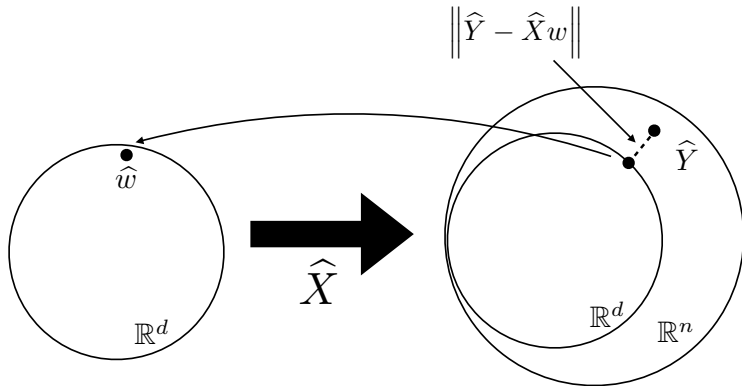
$$\frac{1}{n} \sum_{i=1}^n (y_i - w^\top x_i)^2 = \frac{1}{n} \left\| \hat{Y} - \hat{X}w \right\|^2.$$

This is the least squares problem associated to the linear system

$$\hat{X}w = \hat{Y}.$$

Overdetermined lin. syst. aka underparameterized models

$$n > d$$



Least squares solutions

From the optimality conditions

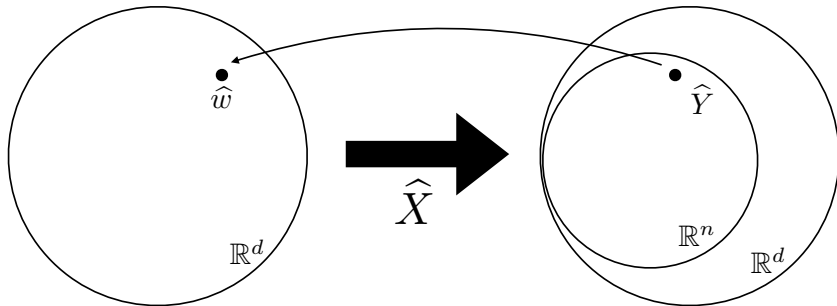
$$\nabla_w \frac{1}{n} \left\| \hat{Y} - \hat{X}w \right\|^2 = 0$$

we can derive the **normal equation**

$$\hat{X}^\top \hat{X}w = \hat{X}^\top \hat{Y} \quad \Leftrightarrow \quad \hat{w} = (\hat{X}^\top \hat{X})^{-1} \hat{X}^\top \hat{Y}.$$

Underdetermined lin. syst.
aka overparameterized model

$$n < d$$



Minimal norm solution

There can be many solutions

$$\hat{X}\hat{w} = \hat{Y}, \quad \text{and} \quad \hat{X}w_0 = 0 \quad \Rightarrow \quad \hat{X}(\hat{w} + w_0) = \hat{Y}.$$

Consider

$$\min_{w \in \mathbb{R}^d} \|w\|^2, \quad \text{subj. to} \quad \hat{X}w = \hat{Y}.$$

Using the method of Lagrange multipliers, the solution is

$$\hat{w} = \hat{X}^\top (\hat{X}\hat{X}^\top)^{-1} \hat{Y}.$$

Computing the minimal norm solution

$$\min_{w \in \mathbb{R}^d} \|w\|^2 \quad \text{subj.t.} \quad \hat{X}w = \hat{Y}.$$

► Lagrangian

$$\mathcal{L}(w, \lambda) = \|w\|^2 + \lambda^\top (\hat{Y} - \hat{X}w).$$

► First-order condition

$$\nabla_w \mathcal{L} = 2w - \hat{X}^\top \lambda = 0 \quad \Rightarrow \quad w = \frac{1}{2} \hat{X}^\top \lambda.$$

► Substitute into the constraint

$$\hat{X}w = \frac{1}{2} \hat{X} \hat{X}^\top \lambda = \hat{Y} \quad \Rightarrow \quad (\hat{X} \hat{X}^\top) \lambda = 2\hat{Y}.$$

► Finally

$$\lambda = 2(\hat{X} \hat{X}^\top)^{-1} \hat{Y}, \quad \hat{w} = \hat{X}^\top (\hat{X} \hat{X}^\top)^{-1} \hat{Y}.$$

Geometric view on minimal norm

$$\hat{w} = \hat{X}^\top (\hat{X} \hat{X}^\top)^{-1} \hat{Y} \quad \Rightarrow \quad \hat{w} \in \text{Ran}(\hat{X}^\top)$$

Recall the orthogonal decomposition

$$\mathbb{R}^d = \text{Ran}(\hat{X}^\top) \oplus \text{Ker}(\hat{X}).$$

Any solution of $\hat{X}w = \hat{Y}$ can be written as

$$w = \hat{w} + w_0, \quad w_0 \in \text{Ker}(\hat{X}),$$

that is the minimal norm solution $\hat{w}$ is the unique one orthogonal to $\text{ker}(\hat{X})$,

Pseudoinverse

$$\hat{w} = \hat{X}^\dagger \hat{Y}$$

For $n > d$, (independent columns)

$$\hat{X}^\dagger = (\hat{X}^\top \hat{X})^{-1} \hat{X}^\top.$$

For $n < d$, (independent rows)

$$\hat{X}^\dagger = \hat{X}^\top (\hat{X} \hat{X}^\top)^{-1}.$$

Spectral view

Consider the SVD of $\hat{X}$

$$\hat{X} = USV^\top \quad \Leftrightarrow \quad \hat{X}w = \sum_{j=1}^r s_j (v_j^\top w) u_j,$$

here $r \leq n \wedge d$ is the rank of $\hat{X}$.

Then,

$$\hat{w}^\dagger = \hat{X}^\dagger \hat{Y} = \sum_{j=1}^r \frac{1}{s_j} (u_j^\top \hat{Y}) v_j.$$

Pseudoinverse and bias

$$\hat{w}^\dagger = \hat{X}^\dagger \hat{Y} = \sum_{j=1}^r \frac{1}{s_j} (u_j^\top \hat{Y}) v_j.$$

$(v_j)_j$ are principal components of $\hat{X}$: OLS “likes” principal components.

Not all linear functions are the same for OLS!

The pseudoinverse introduces a bias towards certain solutions.

Outline

Overparameterization and bias

Regularization and stability

Representer theorem and numerical complexity

Beyond ℓ^2 regularization

From OLS to ridge regression

Recall, it also holds,

$$\hat{X}^\dagger = \lim_{\lambda \rightarrow 0_+} (\hat{X}^\top \hat{X} + \lambda I)^{-1} \hat{X}^\top = \lim_{\lambda \rightarrow 0_+} \hat{X}^\top (\hat{X} \hat{X}^\top + \lambda I)^{-1}.$$

Consider for $\lambda > 0$,

$$\hat{w}_\lambda = (\hat{X}^\top \hat{X} + \lambda I)^{-1} \hat{X}^\top \hat{Y}.$$

This is called ridge regression.

Spectral view on ridge regression

$$\hat{w}_\lambda = (\hat{X}^\top \hat{X} + \lambda I)^{-1} \hat{X}^\top \hat{Y}$$

Considering the SVD of $\hat{X}$,

$$\hat{w}_\lambda = \sum_{j=1}^r \frac{s_j}{s_j^2 + \lambda} (u_j^\top \hat{Y}) v_j.$$

Ridge regression as filtering

$$\hat{w}_\lambda = \sum_{j=1}^r \frac{s_j}{s_j^2 + \lambda} (u_j^\top \hat{Y}) v_j$$

The function

$$F(s) = \frac{s}{s^2 + \lambda},$$

acts as a low pass filter (low frequencies= principal components).

- ▶ For s small, $F(s) \approx 1/\lambda$.
- ▶ For s big, $F(s) \approx 1/s$.

Ridge regression as ERM

$$\hat{w}_\lambda = (\hat{X}^\top \hat{X} + \lambda I)^{-1} \hat{X}^\top \hat{Y}$$

is the solution of

$$\min_{w \in \mathbb{R}^d} \underbrace{\left\| \hat{Y} - \hat{X} w \right\|^2}_{\hat{L}_\lambda(w)} + \lambda \|w\|^2.$$

It follows from,

$$\Delta \hat{L}_\lambda(w) = -\frac{2}{n} \hat{X}^\top (\hat{Y} - \hat{X} w) + 2\lambda w = 2\left(\frac{1}{n} \hat{X}^\top \hat{X} + \lambda I\right) w - \frac{2}{n} \hat{X}^\top \hat{Y}.$$

Ridge regression as ERM

ERM interpretation suggests the rescaling

$$\hat{w}_\lambda = (\hat{X}^\top \hat{X} + \textcolor{red}{n}\lambda I)^{-1} \hat{X}^\top \hat{Y}$$

since

$$\min_{w \in \mathbb{R}^d} \underbrace{\frac{\textcolor{red}{1}}{\textcolor{red}{n}} \|\hat{Y} - \hat{X}w\|^2 + \lambda \|w\|^2}_{\hat{L}_\lambda(w)}.$$

Related ideas

Tikhonov

$$\min_{w \in \mathbb{R}^d} \frac{1}{n} \left\| \hat{Y} - \hat{X}w \right\|^2 + \lambda \|w\|^2$$

Morozov

$$\min_{w \in \mathbb{R}^d} \|w\|^2 \quad \text{subj. to} \quad \frac{1}{n} \left\| \hat{Y} - \hat{X}w \right\|^2 \leq \delta$$

Ivanov

$$\min_{w \in \mathbb{R}^d} \frac{1}{n} \left\| \hat{Y} - \hat{X}w \right\|^2, \quad \text{subj. to} \quad \|w\|^2 \leq R$$

Note: equivalent optimization problems for appropriate parameter choices.

Ridge regression and structural risk minimization

The constraint

$$\|w\|^2 \leq R$$

- ▶ restricts the search of solution,
- ▶ shrinks the solution coefficients.

Different views on regularization

$$\hat{w} = \hat{X}^\dagger \hat{Y}$$

$$\hat{w}_\lambda = (\hat{X}^\top \hat{X} + \lambda I)^{-1} \hat{X}^\top \hat{Y}$$

$$\min_{w \in \mathbb{R}^d} \text{ s.t. } \hat{X}w = \hat{Y} \quad \|w\|^2$$

$$\min_{w \in \mathbb{R}^d} \frac{1}{n} \sum_{i=1}^n (y_i - w^\top x_i)^2 + \lambda \|w\|^2$$

- Introduces a *bias* towards certain solutions: small norm/principal components,
- controls the *stability* of the solution.

Outline

Overparameterization and bias

Regularization and stability

Representer theorem and numerical complexity

Beyond ℓ^2 regularization

Complexity of ridge regression

Back to computations.

Solving

$$\hat{w}^\lambda = (\hat{X}^\top \hat{X} + \lambda I)^{-1} \hat{X}^\top \hat{Y}$$

requires essentially (using a direct solver)

- ▶ time $O(nd^2 + d^3)$,
- ▶ memory $O(nd \vee d^2)$.

What if $n \ll d$?

Representer theorem

A simple observation

We can show that (e.g. using SVD)

$$(\hat{X}^\top \hat{X} + \lambda I)^{-1} \hat{X}^\top = \hat{X}^\top (\hat{X} \hat{X}^\top + \lambda I)^{-1}$$

More on computational complexity

Then

$$\hat{w}_\lambda = \hat{X}^\top (\hat{X} \hat{X}^\top + \lambda I)^{-1} \hat{Y}.$$

requires essentially (using a direct solver)

- ▶ time $O(n^2 d + n^3)$,
- ▶ memory $O(nd \vee n^2)$.

Representer theorem

Note that

$$\hat{w}_\lambda = \hat{X}^\top \underbrace{(\hat{X}\hat{X}^\top + \lambda I)^{-1}}_{c \in \mathbb{R}^n} \hat{Y} = \sum_{i=1}^n x_i c_i.$$

The coefficients vector is a linear combination of the input points.

Then

$$\hat{f}_\lambda(x) = x^\top \hat{w}_\lambda = x^\top \hat{X}^\top c = \sum_{i=1}^n x^\top x_i c_i$$

The function we obtain is a linear combination of *inner products*.

This will be the key to nonparametric (kernel) learning.

Outline

Overparameterization and bias

Regularization and stability

Representer theorem and numerical complexity

Beyond ℓ^2 regularization

Other regularization terms

$$\frac{1}{n} \sum_{i=1}^n (y_i - w^\top x_i)^2 + \lambda R(w)$$

Beyond ℓ^2 , other regularization terms can be considered.

Popular choices are sparsity-inducing regularizers, e.g.

► ℓ^0 “norm”:

$$R(w) = \|w\|_0 = \#\{j : w_j \neq 0\},$$

► ℓ^1 norm:

$$R(w) = \|w\|_1 = \sum_j |w_j|.$$

Why sparsity?

$$\min_{w \in \mathbb{R}^d} \frac{1}{n} \|\hat{Y} - \hat{X}w\|^2 + \lambda \|w\|_0$$

Advantages:

- Variable selection
- Interpretability

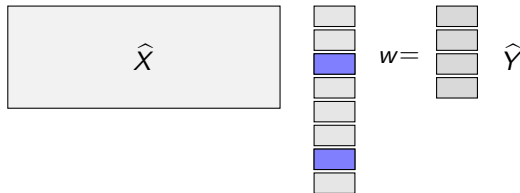
Challenges:

- Nonconvex optimization
- NP-hard in general

Sparsity and overparametrization

$$\min_{w \in \mathbb{R}^d} \frac{1}{n} \|\hat{Y} - \hat{X}w\|^2 + \lambda \|w\|_0$$

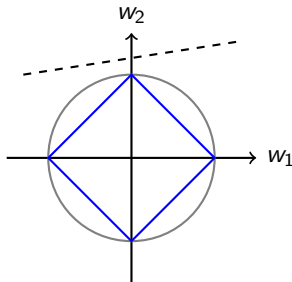
$$\begin{aligned} \min_{w \in \mathbb{R}^d} & \|w\|_0 \\ \text{s.t.} & \hat{X}w = \hat{Y} \end{aligned}$$



Why ℓ^1

$$\|w\|_0 \mapsto \|w\|_1 \neq \|w\|$$

$$\begin{aligned} \min_{w \in \mathbb{R}^d} \|w\|_1 \\ \text{s.t. } \hat{X}w = \hat{Y} \end{aligned}$$



Computing the Lasso/Basis pursuit

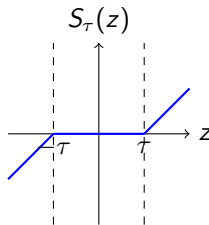
$$\min_{w \in \mathbb{R}^d} \frac{1}{n} \|\hat{Y} - \hat{X}w\|^2 + \lambda \|w\|_1$$

ISTA: Iterative soft-thresholding

$$w^{(k+1)} = S_{\lambda\eta} \left(w^{(k)} - \eta \hat{X}^\top (\hat{X} w^{(k)} - \hat{Y}) \right)$$

Soft-thresholding

$$S_\tau(z) = \text{sign}(z) \max\{|z| - \tau, 0\}$$



Remarks

- ▶ Other sparse regularizers: group lasso, TV-norm ...
- ▶ Other optimizers: Frank-Wolfe/conditional gradient (matching pursuit).
- ▶ No representer theorem.

Summing up

- ▶ ERM.
- ▶ From OLS to ridge regression.
- ▶ Different views: (spectral) filtering and ERM.
- ▶ Regularization and bias.
- ▶ Sparse regularization.